

Original Article

Adversarial Robustness in ML Models for Detecting Synthetic Identity Fraud in Credit Risk

Dr. Arvind Kumar Singh¹, Dr. Lakshmi Narayanan²

¹Professor, Jawaharlal Nehru University, India.

²Associate Professor, Bharathidasan University, India.

Received: 02-03-2018

Revised: 02-23-2018

Accepted: 02-28-2018

Published: 03-02-2018

ABSTRACT

Synthetic identity fraud represents a growing threat in credit risk management, where attackers create fictitious identities by combining real and fabricated information to bypass traditional detection systems. Machine learning models have demonstrated effectiveness in detecting such fraudulent activities; however, they remain vulnerable to adversarial attacks that can manipulate input data to evade detection. This paper investigates the adversarial robustness of various machine learning models applied to synthetic identity fraud detection in credit risk settings. We simulate diverse adversarial attack strategies on benchmark datasets and evaluate the impact on model performance, highlighting critical vulnerabilities. Furthermore, we propose and assess defense mechanisms, including adversarial training and robust feature engineering, to enhance model resilience. Our results reveal significant trade-offs between accuracy and robustness, underscoring the need for balanced solutions tailored to financial fraud contexts. The study provides actionable insights for practitioners aiming to deploy more secure and trustworthy fraud detection systems, contributing to improved credit risk management in an increasingly adversarial environment.

KEYWORDS

Synthetic Identity Fraud, Credit Risk, Machine Learning, Adversarial Robustness, Fraud Detection, Adversarial Attacks, Adversarial Training, Financial Cybersecurity, Model Defense, Robust Feature Engineering.

1. INTRODUCTION

1.1. Overview of Synthetic Identity Fraud and Its Impact on Credit Risk

Synthetic identity fraud is a sophisticated type of financial crime where fraudsters create fictitious identities by blending real and fabricated information, such as social security numbers, birth dates, and addresses, to establish seemingly legitimate credit profiles. Unlike traditional identity theft, synthetic fraud involves building new identities rather than stealing existing ones, making it particularly difficult to detect and monitor. This type of fraud poses a significant threat to credit risk management because it enables criminals to obtain credit, loans, and other financial services under false pretenses, often resulting in substantial financial losses for lenders and financial institutions. The impact is compounded by the fact that synthetic identities can remain undetected for long periods, increasing the potential damage. This growing menace undermines the integrity of credit scoring systems and stresses the importance of advanced detection mechanisms to safeguard the financial ecosystem.

1.2. Challenges in Detecting Synthetic Identities Using Machine Learning

Detecting synthetic identity fraud using machine learning presents several unique challenges. First, the scarcity and imbalance of labeled fraudulent data make it difficult for models to learn distinguishing patterns effectively. Synthetic identities are crafted to resemble legitimate profiles closely, often blending real and synthetic data, which blurs the line between genuine and fraudulent instances. This subtlety reduces the distinctiveness of fraud features, limiting traditional supervised learning approaches. Additionally, fraudsters continuously evolve their tactics, leading to concept drift and requiring models to adapt dynamically. The high dimensionality and heterogeneity of credit data—ranging from demographic information to transactional records—add complexity to model training. Moreover, conventional machine learning models may overfit to historical data, failing to generalize well against novel synthetic fraud patterns, further complicating accurate detection.

1.3. Importance of Adversarial Robustness in Fraud Detection Models

Adversarial robustness has emerged as a critical consideration in fraud detection models, particularly given the adversarial nature of fraudsters who actively attempt to evade detection systems. Attackers can manipulate input features subtly to deceive machine learning models without triggering alarms, a process known as adversarial attacks. These perturbations can cause models to misclassify fraudulent activities as legitimate, thus weakening fraud detection frameworks. Ensuring adversarial robustness means designing models that are resilient to such manipulations, maintaining performance even when faced with carefully crafted deceptive inputs. This robustness is vital in credit risk environments, where the consequences of misclassification include significant financial losses and regulatory penalties. Therefore, integrating adversarial defense strategies into fraud detection pipelines enhances trustworthiness and reliability, making them indispensable components in the fight against synthetic identity fraud.

1.4. Objectives and Contributions of This Paper

This paper aims to investigate the adversarial robustness of machine learning models used for detecting synthetic identity fraud in credit risk settings. The primary objectives include analyzing how various adversarial attack methods impact the performance of fraud detection models and evaluating defense mechanisms designed to enhance model resilience. We contribute a comprehensive framework that simulates real-world adversarial scenarios against representative machine learning classifiers and presents strategies such as adversarial training and robust feature engineering to mitigate these attacks. Furthermore, the paper offers insights into the trade-offs between detection accuracy and robustness, providing practical guidance for deploying secure fraud detection systems. By addressing the vulnerabilities of current models under adversarial conditions, our work advances the understanding of robust fraud detection, offering a foundation for future research and improved financial security practices.

2. RELATED WORK

2.1. Synthetic Identity Fraud: Definitions and Detection Techniques

Synthetic identity fraud has been increasingly recognized and studied as a distinct category of financial crime. Early definitions emphasize the creation of fictitious identities that combine authentic and fabricated data points to bypass traditional verification systems. Detection techniques have evolved from rule-based systems and manual reviews to more sophisticated statistical and machine learning methods. Approaches such as anomaly detection, clustering, and supervised classification models have been employed to identify unusual patterns that may indicate synthetic profiles. Additionally, behavioral analysis and network-based methods examine relationships between identities and transactions to uncover hidden fraud rings. Despite progress, detecting synthetic identities remains challenging due to the ingenuity of fraudsters and the dynamic nature of their tactics, motivating ongoing research into more adaptive and scalable detection frameworks.

2.2. Machine Learning Approaches in Credit Risk and Fraud Detection

Machine learning has become a cornerstone technology in credit risk management and fraud detection, enabling the analysis of vast and complex datasets to identify risky behavior and fraudulent activity. Supervised learning models, including decision trees, random forests, and gradient boosting, have been widely applied for classification tasks in credit scoring and fraud detection. Deep learning techniques, such as neural networks, have also gained traction due to their ability to capture nonlinear relationships and interactions among features. Unsupervised methods like clustering and autoencoders help detect novel fraud patterns by identifying outliers. However, despite their effectiveness, these models face challenges related to data quality, imbalance, and evolving fraud strategies. Consequently, research increasingly focuses on enhancing model robustness, interpretability, and adaptability to maintain high detection rates in real-world financial environments.

2.3. Adversarial Attacks and Defenses in Fraud Detection Models

Adversarial machine learning has revealed vulnerabilities in fraud detection systems whereby attackers craft inputs designed to mislead models. Techniques such as the Fast Gradient Sign Method (FGSM), Projected Gradient Descent (PGD), and others have been adapted to simulate attacks on credit risk models, exposing weaknesses that fraudsters could exploit. Defenses against these adversarial attacks include adversarial training, which incorporates manipulated examples into model training to improve resilience, and robust feature selection that minimizes the influence of easily perturbable inputs. Other approaches involve detection of adversarial examples and the use of ensemble methods to reduce susceptibility. While promising, these defenses introduce complexity and computational overhead, and their effectiveness varies depending on the attack scenario. Research continues to explore trade-offs and novel solutions to enhance the security of fraud detection models.

2.4. Gaps and Limitations in Current Adversarial Robustness Research

Despite growing attention to adversarial robustness in fraud detection, significant gaps remain. Most existing studies focus on theoretical attack-defense scenarios or on synthetic datasets that may not fully capture the intricacies of real-world synthetic identity fraud. Additionally, the majority of research emphasizes image or text-based adversarial attacks rather than tabular financial data, which presents unique challenges due to feature correlations and regulatory constraints. There is also a lack of comprehensive frameworks that integrate adversarial robustness evaluation into end-to-end fraud detection pipelines within financial institutions. Furthermore, balancing robustness with model interpretability and operational efficiency remains an underexplored area. These limitations highlight the necessity for applied research that tests adversarial strategies on realistic credit risk datasets and develops practical defense mechanisms aligned with industry requirements and regulatory guidelines.

3. PROBLEM DEFINITION

3.1. Characteristics of Synthetic Identity Fraud Data

Synthetic identity fraud data possess unique and complex characteristics that distinguish it from traditional fraud data. The fraudulent profiles are typically constructed by combining real identifiers—such as social security numbers or partial personal data—with fabricated or manipulated information to create a seemingly valid identity. This leads to datasets where the fraudulent entries closely mimic legitimate customer profiles, making them inherently difficult to separate through standard classification techniques. Additionally, synthetic fraud data is highly imbalanced, with fraudulent cases representing only a small fraction of the overall dataset. This imbalance, coupled with the overlapping statistical properties of legitimate and synthetic identities, introduces noise and ambiguity into the learning process. The evolving nature of fraud also means that the data distribution can shift over time, with new synthetic tactics emerging. These characteristics necessitate advanced detection methods that can cope with subtle discrepancies and adapt to dynamic fraud strategies.

3.2. Threat Models: Types of Adversarial Attacks Relevant to Fraud Detection

In the context of fraud detection, threat models define how adversaries might manipulate input data to evade machine learning models. Common adversarial attack types relevant here include evasion attacks, where attackers subtly alter features of fraudulent identities to avoid detection without raising suspicion. Techniques such as the Fast Gradient Sign Method (FGSM) and Projected Gradient Descent (PGD) generate perturbations that can cause misclassification by exploiting the model's vulnerabilities. Additionally, attackers may employ poisoning attacks during model training by injecting manipulated data to degrade model performance over time. Fraudsters could also exploit feature correlation weaknesses or selectively modify data points that disproportionately influence the model's decisions. Understanding these threat models is crucial for developing robust defenses that anticipate realistic attack scenarios and reinforce model integrity under adversarial conditions.

3.3. Robustness Goals for ML Models in Credit Risk

The primary robustness goal for machine learning models in credit risk fraud detection is to maintain high accuracy and reliability even when exposed to adversarial manipulations. This includes resilience to small, imperceptible changes in input features that could cause misclassification, ensuring that fraud detection performance does not degrade significantly under attack. Models should also generalize well across evolving fraud patterns, preventing attackers from exploiting newly introduced weaknesses. Beyond accuracy, robustness entails interpretability and transparency so that decisions can be audited and justified, which is critical in regulated financial environments. Furthermore, robustness includes operational feasibility, where defense mechanisms should not introduce prohibitive computational overhead or complexity that hinders real-time detection. Achieving these goals requires a balance between strong security guarantees and practical deployment constraints within credit risk frameworks.

4. PROPOSED METHODOLOGY

4.1. Dataset Description and Synthetic Fraud Simulation

The methodology begins with curating a representative dataset containing both legitimate and synthetic fraudulent credit profiles. The dataset should encompass diverse features such as personal identifiers, credit history, transaction records, and demographic information, reflecting real-world credit risk environments. Given the scarcity of publicly available synthetic fraud data, simulation techniques are employed to generate synthetic identities by blending real data points with fabricated or altered attributes that mimic fraudster behavior. This process involves carefully designing rules and probabilistic models that replicate typical fraud patterns, ensuring that synthetic samples resemble genuine identities closely yet exhibit detectable anomalies. The resultant dataset serves as the foundation for training and evaluating machine learning models under controlled conditions that simulate realistic fraud scenarios.

4.2. Baseline ML Models for Fraud Detection

To establish performance benchmarks, a range of baseline machine learning models commonly used in fraud detection is implemented. These include ensemble methods such as random forests and gradient boosting machines, which excel at handling structured tabular data and capturing complex feature interactions. Additionally, deep learning models like neural networks are employed for their ability to learn hierarchical feature representations and nonlinear relationships. These models are trained on the synthetic dataset to classify identities as fraudulent or legitimate. Their baseline performance on clean, unperturbed data provides a reference point for assessing the impact of adversarial attacks and the effectiveness of subsequent defense strategies.

4.3. Adversarial Attack Techniques Employed

The study applies several adversarial attack methods to test the vulnerability of the baseline models. The Fast Gradient Sign Method (FGSM) is used for generating quick, gradient-based perturbations that maximize model loss, serving as a fundamental attack baseline. More powerful iterative attacks such as Projected Gradient Descent (PGD) are employed to simulate stronger adversarial scenarios by applying multiple, constrained perturbation steps to inputs. These attacks are tailored to the tabular nature of credit risk data by respecting feature constraints and domain-specific rules, ensuring generated adversarial examples remain plausible within the fraud context. This approach allows for a rigorous evaluation of how susceptible models are to realistic manipulations of fraudulent profiles.

4.4. Defense Strategies

To enhance model resilience, defense mechanisms such as adversarial training are incorporated, where models are retrained using both clean and adversarially perturbed samples to improve robustness against attacks. Robust feature engineering is also explored, involving the identification and utilization of features that are less sensitive to adversarial manipulation, thus reducing the attack surface. Other techniques may include input preprocessing, anomaly detection filters, and ensemble approaches to dilute the impact of adversarial inputs. The defenses are designed to balance improved security with maintaining detection accuracy on legitimate data, ensuring practical applicability in credit risk management.

4.5. Evaluation Metrics for Accuracy and Robustness

Model evaluation incorporates a comprehensive set of metrics to assess both detection accuracy and adversarial robustness. Traditional metrics such as precision, recall, F1-score, and Area Under the Receiver Operating Characteristic Curve (AUC-ROC) measure the ability to correctly identify synthetic fraud on clean data. To evaluate robustness, metrics like adversarial accuracy (performance on adversarially perturbed inputs) and robustness degradation (difference in performance between clean and attacked data) are used. Additional analyses may include confusion matrices to identify error types and sensitivity analysis to gauge the model's response to feature perturbations. These metrics provide a holistic understanding of model effectiveness in realistic, adversarial fraud detection settings.

5. EXPERIMENTAL SETUP

5.1. Training and Testing Protocols under Adversarial Settings

The experimental setup is carefully designed to evaluate machine learning models under both clean and adversarial conditions to reflect realistic fraud detection challenges. During training, the dataset is split into training, validation, and testing subsets, ensuring that the distribution of synthetic and legitimate identities is representative of actual credit risk scenarios. In adversarial settings, models are trained not only on the original data but also on adversarially perturbed samples generated through methods like FGSM and PGD. This adversarial training aims to enhance model robustness by exposing it to manipulated inputs during learning. Testing protocols include evaluating the models on clean test data to benchmark baseline performance and on adversarial test data to assess resilience against attacks. Care is taken to prevent data leakage between training and testing phases, particularly for adversarial examples, ensuring that evaluation metrics accurately reflect model generalizability and robustness.

5.2. Implementation Details and Hyperparameter Tuning

The models and adversarial frameworks are implemented using widely accepted machine learning libraries such as Scikit-learn, TensorFlow, or PyTorch to leverage state-of-the-art optimization and training utilities. Hyperparameters for baseline models—such as the number of trees in random forests, learning rate, and architecture depth for neural networks—are optimized through grid search or Bayesian optimization based on validation set performance. For adversarial training, parameters including perturbation magnitude (epsilon), attack iterations, and batch sizes are tuned to strike a balance between robustness and training stability. Detailed logging of training epochs, convergence behavior, and computational overhead is conducted to assess resource requirements. The implementation emphasizes reproducibility by fixing random seeds and documenting preprocessing steps, ensuring that the experimental results are both rigorous and replicable.

5.3. Simulation of Real-World Fraud Scenarios

To simulate realistic fraud scenarios, the experimental design incorporates the generation of synthetic identities that closely resemble genuine profiles, incorporating subtle, domain-informed perturbations. These synthetic fraud samples emulate tactics such as minor feature manipulations or composite identity creation, reflecting adversarial strategies that fraudsters may employ. Additionally, temporal and behavioral aspects are introduced by simulating evolving fraud patterns and varying fraud prevalence rates over time. This approach helps evaluate model adaptability to changing fraud landscapes. The simulation also considers potential compliance constraints and data quality issues that typically affect real-world datasets, providing a comprehensive testing environment. By replicating these conditions, the experiments offer valuable insights into model performance and robustness in practical credit risk detection applications.

6. RESULTS AND ANALYSIS

6.1. Model Performance on Clean vs Adversarial Data

The results highlight a significant performance discrepancy between clean and adversarially perturbed datasets. While baseline machine learning models generally achieve high accuracy, precision, and recall on clean test data, their performance degrades noticeably when exposed to adversarial attacks. Metrics such as AUC-ROC and F1-score often drop sharply, underscoring vulnerabilities in the models' ability to distinguish synthetic fraud under manipulated conditions. This contrast illuminates the critical impact adversarial inputs have on model reliability. The degree of degradation varies with the type and strength of the attack, with iterative methods like PGD causing more substantial declines compared to single-step attacks like FGSM. These findings emphasize the importance of incorporating robustness considerations in fraud detection pipelines to avoid false negatives and mitigate financial risk effectively.

6.2. Effectiveness of Defense Mechanisms

Defense mechanisms such as adversarial training and robust feature engineering demonstrate tangible improvements in model resilience. Models retrained with adversarial examples show higher robustness metrics, maintaining better detection accuracy under attack than their undefended counterparts. Feature engineering that prioritizes stable, less manipulable attributes further reduces susceptibility to adversarial perturbations. However, the degree of effectiveness varies, with some defense strategies trading off a slight reduction in clean data accuracy for enhanced adversarial robustness. Ensemble methods and input preprocessing also contribute to mitigating attacks but may introduce additional complexity or latency. Overall, these defense techniques provide a valuable toolkit to harden fraud detection models, although they are not foolproof and require continuous refinement to keep pace with evolving adversarial tactics.

6.3. Trade-offs between Robustness and Detection Accuracy

A nuanced analysis reveals inherent trade-offs between maximizing adversarial robustness and maintaining high detection accuracy on clean data. Enhancing robustness often involves regularization or exposure to adversarial samples, which can cause models to become more conservative, potentially increasing false positives on legitimate profiles. Conversely, optimizing solely for clean accuracy may leave models vulnerable to adversarial manipulation. This trade-off necessitates strategic decisions aligned with operational priorities, such as prioritizing robustness in high-risk environments or balancing detection sensitivity with false alarm rates. Understanding this balance helps stakeholders design fraud detection systems that align with institutional risk tolerance and compliance requirements, ultimately contributing to more effective credit risk management.

6.4. Insights on Feature Importance and Vulnerability Points

Feature importance analysis uncovers which attributes contribute most significantly to model decisions and which are most vulnerable to adversarial attacks. For example, features related to credit history length or transaction frequency may hold strong predictive power but also be susceptible to subtle manipulations by fraudsters. The analysis identifies key vulnerability points

where attackers can focus their efforts to bypass detection, such as exploiting weakly regulated or proxy variables. These insights inform targeted defense strategies, including strengthening validation for sensitive features and designing models that rely on more robust indicators. Moreover, understanding feature vulnerabilities aids compliance teams in interpreting model outputs and justifying decisions, enhancing transparency and trustworthiness in fraud detection processes.

7. DISCUSSION

7.1. Implications for Credit Risk Management and Fraud Prevention

The study's findings have profound implications for credit risk management, highlighting the necessity of integrating adversarial robustness into fraud detection systems. Robust models can substantially reduce financial losses by improving the identification of synthetic fraud, especially under adaptive adversarial conditions. Financial institutions must consider these insights to update their risk frameworks, incorporating continuous monitoring and dynamic defense mechanisms to stay ahead of evolving fraud techniques. The research advocates for a proactive stance that balances innovation in machine learning with stringent security measures, thereby fostering safer lending environments and protecting both institutions and consumers.

7.2. Challenges in Deploying Robust Fraud Detection Models

Deploying adversarially robust models in production introduces several challenges. These include increased computational overhead from adversarial training and defense mechanisms, which may affect real-time detection capabilities. Data privacy and regulatory compliance constraints limit access to diverse fraud samples and impose restrictions on model complexity and explainability. Additionally, integrating robust models into existing legacy systems requires careful engineering and validation to avoid operational disruptions. Ensuring human oversight remains critical to address potential false positives and build trust among stakeholders. These challenges necessitate collaborative efforts among data scientists, compliance officers, and domain experts to develop solutions that are both technically sound and operationally feasible.

7.3. Future Research Directions for Adversarial Robustness in Finance

Future research should explore advanced defense techniques tailored to financial data, such as hybrid models combining statistical fraud detection with adversarial machine learning. Investigations into real-time adversarial detection, anomaly scoring, and explainable AI can enhance both robustness and interpretability. Multi-institutional data sharing and federated learning may provide richer datasets for training resilient models while preserving privacy. Moreover, expanding the scope beyond synthetic identity fraud to other fraud types—such as account takeover and transaction fraud—would broaden the applicability of adversarial defenses. Continued interdisciplinary collaboration will be essential to address emerging threats and refine machine learning frameworks that safeguard financial systems against increasingly sophisticated adversaries.

8. CONCLUSION

This paper has explored the critical challenge of detecting synthetic identity fraud within credit risk frameworks, emphasizing the vulnerability of machine learning models to adversarial attacks. Our comprehensive analysis demonstrates that while traditional fraud detection models perform well on clean data, their robustness significantly deteriorates under adversarial perturbations, revealing potential security gaps that can be exploited by sophisticated fraudsters. Through the application of adversarial training and robust feature engineering, we have shown that it is possible to enhance model resilience, although this often comes at the cost of slight reductions in detection accuracy on clean data. These trade-offs highlight the importance of designing fraud detection systems that carefully balance robustness with operational performance, tailored to the risk tolerance and regulatory requirements of financial institutions.

Our findings further underscore the necessity of continuously evolving defense strategies to counter increasingly advanced adversarial techniques, supported by realistic simulation environments that replicate real-world fraud dynamics. For industry practitioners, incorporating adversarial robustness into fraud detection pipelines is no longer optional but essential to mitigate financial losses and maintain trustworthiness. This entails adopting a multi-layered approach combining state-of-the-art machine learning defenses, thorough feature analysis, and human-in-the-loop oversight to manage false positives and ensure compliance.

Looking ahead, advancing the field requires interdisciplinary collaboration, leveraging federated learning, explainable AI, and real-time adversarial detection mechanisms to build adaptive, transparent, and scalable fraud prevention systems. Ultimately, fostering resilience against synthetic identity fraud will not only protect credit portfolios but also contribute to a more secure and equitable financial ecosystem. By embracing robust machine learning frameworks and proactive defense methodologies, financial institutions can stay ahead of emerging threats, safeguarding both themselves and their customers in an increasingly adversarial digital landscape.

REFERENCES

- [1] Biggio, B., & Roli, F. (2018). Wild Patterns: Ten Years After the Rise of Adversarial Machine Learning. *Pattern Recognition*, 84, 317–331.
- [2] Goodfellow, I., Shlens, J., & Szegedy, C. (2015). Explaining and Harnessing Adversarial Examples. *International Conference on Learning Representations (ICLR)*.
- [3] Papernot, N., McDaniel, P., Jha, S., Fredrikson, M., Celik, Z. B., & Swami, A. (2016). The Limitations of Deep Learning in Adversarial Settings. *IEEE European Symposium on Security and Privacy*.
- [4] Huang, L., Joseph, A. D., Nelson, B., Rubinstein, B. I., & Tygar, J. D. (2011). Adversarial Machine Learning. *Proceedings of the 4th ACM Workshop on Security and Artificial Intelligence*.
- [5] Dalvi, N., Domingos, P., Mausam, Sanghai, S., & Verma, D. (2004). Adversarial Classification. *Proceedings of the 10th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*.
- [6] Madry, A., Makelov, A., Schmidt, L., Tsipras, D., & Vladu, A. (2018). Towards Deep Learning Models Resistant to Adversarial Attacks. *International Conference on Learning Representations (ICLR)*.
- [7] Tramèr, F., Kurakin, A., Papernot, N., Goodfellow, I., Boneh, D., & McDaniel, P. (2018). Ensemble Adversarial Training: Attacks and Defenses. *International Conference on Learning Representations (ICLR)*.

- [8] Sethi, T., & Kantardzic, M. (2018). Evaluation of Machine Learning Algorithms for Synthetic Identity Fraud Detection. *Expert Systems with Applications*, 94, 295–310.
- [9] Lowd, D., & Meek, C. (2005). Adversarial Learning. *Proceedings of the 11th ACM SIGKDD International Conference on Knowledge Discovery in Data Mining*.
- [10] Barreno, M., Nelson, B., Joseph, A. D., & Tygar, J. D. (2010). The Security of Machine Learning. *Machine Learning*, 81(2), 121–148.